

AN IMPROVED MULTINOMIAL NAÏVE BAYES ALGORITHM WITH LAPLACE SMOOTHING FOR NEWS CLASSIFICATION: THE BISALLAH NEWS NETWORK MODEL (BNNM)

¹Hashim Ibrahim Bisallah, ²Christopher Ubaka Ebelogu

¹Department of Computer Science, University of Abuja, Nigeria

²Department of Computer Science, Ignatius Ajuru University of Education, Portharcourt, Nigeria

hashim.bisallah@uniabuja.edu.ng, christopher.ebelogu@iaue.edu.ng

Abstract

Naïve Bayes remains one of the most widely used algorithms for text classification and sentiment analysis, yet its application to Nigerian newspaper social-media data has not been previously documented. This study develops and evaluates an improved Multinomial Naïve Bayes classifier - incorporating TF-IDF feature weighting and Laplace (additive) smoothing - for the automatic classification of Nigerian newspaper tweets into politics, sports, and business categories. Drawing on a cleaned corpus of 24,679 tweets from *The Punch*, *Vanguard*, and *The Guardian* (2015–2018), the classifier was trained on a 70/30 train-test split and tuned via grid-search cross-validation across *n*-gram ranges, TF-IDF normalisation schemes, and smoothing parameters (α). The resulting model - termed the Bisallah News Network Model (BNNM) - achieved an overall classification accuracy of 83.1%, with category-level prediction accuracies of 80.6% for politics, 72.1% for sports, and 96.0% for business, an average recall of 0.75, and an average F1-score of 0.77. These results are broadly comparable to, though in some cases below, prior studies (e.g., Preety & Dahiya, 2015, at 89.01%; Perdana & Pinandito, 2018, at 83.8%), with differences attributable to corpus size and composition. The study demonstrates that an improved, Laplace-smoothed Multinomial Naïve Bayes pipeline can deliver practically useful classification performance on noisy, real-world Nigerian social-media text, offering a foundation for automated content categorisation and editorial decision-support tools in African digital newsrooms.

Keywords: Naïve Bayes; Laplace smoothing; Text Classification; Big Data; Nigerian Digital News Media; Machine Learning; TF-IDF

Introduction

The classification of large volumes of unstructured text into meaningful categories is a foundational task in big-data analytics, with direct applications in news media for organising content, understanding audience interests, and informing editorial strategy. Among the algorithms applied to this task, Naïve Bayes has consistently demonstrated a favourable balance of computational efficiency, simplicity, and classification accuracy on large datasets (Kaviani & Dhotre, 2017), and continues to be used as a strong, low-cost baseline even in the era of transformer-based language models (Zhang, 2025).

A substantial body of research has applied Naïve Bayes - often in combination with other techniques - to social-media and news-text classification across diverse contexts: *n*-gram-based event classification (Azam, Jahiruddin & Al-Hasan Haldar, 2015), retweet sentiment modelling (Wang, Zuo & Wang, 2015), opinion classification for customer reviews (Preety & Dahiya, 2015; Khobragade & Jethani, 2017), news-curator classification (Sembodo, Setiawan & Baizal, 2017), fake-news detection (Stahl, 2018), and hate-speech detection (Kiilu et al., 2018), among others. Several of these studies have specifically proposed improvements to baseline Naïve Bayes - through discriminative model components (Miao et al., 2018), Fisher Score feature combination (Perdana & Pinandito, 2018), and feature expansion via hypernym/hyponym relations (Irfani, Ali & Sari, 2018). More recent work continues this trajectory: Zeng et al. (2023) propose an EM-based correction for mislabelled training data that substantially improves Naïve Bayes robustness, while Zhang (2025) demonstrates that a

hybrid TF-IDF/Naïve Bayes framework augmented with domain-specific n-grams and entity-enhanced weighting can yield meaningful F1 gains over conventional implementations while remaining far cheaper to deploy than transformer-based alternatives.

Despite this volume of work, no prior study has applied an improved Naïve Bayes approach to Nigerian (or broader African) newspaper social-media text. This represents a significant gap given the scale and growth of Nigerian newspapers' digital and social-media presence - a gap reinforced by the fact that even the most recent Nigeria-focused Twitter studies (e.g., election-sentiment analyses using BERT, LSTM, and SVM architectures; PMC, 2023) have concentrated on single political events rather than the routine, multi-category news classification problem facing newsrooms day to day. This study addresses that gap by developing and evaluating an improved Multinomial Naïve Bayes classifier - incorporating TF-IDF weighting and Laplace smoothing - for categorising Nigerian newspaper tweets into politics, sports, and business, and by benchmarking its performance against comparable studies from the international literature.

Recent advances have further improved the robustness and applicability of Naïve Bayes through enhanced feature engineering and optimisation techniques. Qianhan Zeng et al. (2023) introduced an expectation-maximisation (EM) approach to address label noise, substantially improving classification performance without requiring extensive manual relabelling (Zeng et al., 2023). Similarly, Li Zhang (2025) demonstrated that integrating TF-IDF weighting with domain-specific n-grams and entity-enhanced feature extraction significantly improves the effectiveness of Naïve Bayes for online news classification while maintaining low computational cost compared with transformer-based models (Zhang, 2025). Furthermore, recent comparative studies indicate that although transformer-based models such as BERT generally achieve superior predictive accuracy, Multinomial Naïve Bayes remains an attractive baseline because of its simplicity, interpretability, computational efficiency, and rapid inference speed, particularly in resource-constrained environments and moderate-sized labelled datasets (Zeng et al., 2023; Zhang, 2025). These findings suggest that enhancing the traditional Naïve Bayes framework through improved preprocessing, TF-IDF weighting, Laplace smoothing, and parameter optimisation remains a viable research direction for efficient news classification and large-scale text mining applications.

Despite the considerable progress achieved in text mining and news classification, limited attention has been given to the application of improved Naïve Bayes techniques for the automatic classification of Nigerian digital newspaper content. Existing studies have largely focused on sentiment analysis, fake news detection, or election-related social media analysis, with little emphasis on routine multi-category news classification within the Nigerian media landscape. This study addresses this gap by developing the Bisallah News Network Model (BNNM), an improved Multinomial Naïve Bayes classifier that integrates TF-IDF feature weighting, Laplace smoothing, and hyperparameter optimisation for the classification of Nigerian newspaper tweets into Politics, Sports, and Business categories. Furthermore, the study introduces a confidence-based prediction threshold that enables the model to reject low-confidence predictions rather than forcing uncertain tweets into inappropriate classes. The performance of the proposed model is evaluated using a corpus of 24,679 tweets collected from three leading Nigerian newspapers and benchmarked against existing Naïve Bayes-based text classification studies. The findings are expected to contribute to the development of efficient, interpretable, and computationally inexpensive intelligent news classification systems for digital newsrooms and other large-scale text mining applications.

Literature Review

Multiple strands of prior research inform the design of the proposed model. Bachhety, Dhingra, Jain, and Jain (2018) proposed an improved Multinomial Naïve Bayes approach for sentiment analysis on social media, demonstrating the value of refining the base algorithm rather than treating it as a fixed baseline. Perdana and Pinandito (2018) combined Naïve Bayes with Fisher Score for textual and non-textual feature integration, achieving 83.8% accuracy with weighting parameters $\alpha = 0.6$ and $\beta = 0.4$ - a benchmark closely comparable to the present study's target performance range. Preety and Dahiya (2015) reported 89.01% accuracy using Naïve Bayes and SVM on a corpus of 1,000 mobile reviews, while Khobragade and Jethani (2017) reported 85.7% accuracy for opinion classification. Alharbi, Alnnamlah, and Liyakathunisa (2018) achieved very high (99.39%) accuracy for sentiment classification of customer tweets using Naïve Bayes via Apache Hadoop and Mahout - though on a markedly different (and likely smaller or more homogeneous) dataset.

More recently, Zeng, Zhu, Zhu, Wang, Zhao, Sun, Su, and Wang (2023) addressed the long-standing problem of mislabelled training data in text classification, proposing an EM-algorithm-based correction to the Naïve Bayes log-likelihood that improved performance without requiring subjective relabelling - a finding directly relevant to noisy, crowd-generated tweet corpora such as the one used here. Zhang (2025) extended the classical TF-IDF/Naïve Bayes pipeline with domain-specific n-grams and entity-enhanced weighting for online news classification, reporting a 7.4% F1 improvement over conventional implementations alongside low inference latency, reinforcing the continued practical relevance of Naïve Bayes-based pipelines for high-volume, real-time newsroom applications. Collectively, these studies suggest that Naïve Bayes performance is highly sensitive to corpus size, label balance, domain, and the specific smoothing and feature-weighting choices applied - factors directly addressed in the present study's methodology.

Methodology

Data Collection

Tweets were collected from the official Twitter (now X) accounts of three leading Nigerian newspapers: The Punch, Vanguard, and The Guardian. Historical tweets posted between January 2015 and December 2018 were retrieved using the Twitter API (Academic Research access) together with manual verification to ensure completeness of the downloaded records. Search parameters included tweet language (English), publication date, source account, and tweet text. Duplicate tweets, retweets, advertisements, and records containing only hyperlinks or multimedia without meaningful textual content were removed during preprocessing. The resulting corpus consisted of 24,679 unique tweets representing three major news categories.

Although the dataset covers tweets posted between 2015 and 2018, it was deliberately selected because it represents a complete, stable, and fully labelled corpus that enabled reproducible experimentation. The study focuses on evaluating methodological improvements to the Multinomial Naïve Bayes algorithm rather than analysing contemporary news events. Since the underlying classification algorithm is domain-independent, the corpus remains suitable for evaluating classification performance while ensuring direct comparison with earlier Naïve Bayes studies that employed similarly archived datasets.

Data and Labelling

The dataset comprised 24,679 cleaned, deduplicated tweets from three Nigerian newspapers' Twitter accounts (2015–2018). Each tweet was assigned to one of three categories (Politics, Sports, and Business) according to the editorial section from which it originated on the respective newspaper account. Where category labels were ambiguous, manual verification was performed using the linked news article and tweet content to ensure consistency. The final

labelled dataset contained 11,312 Politics tweets, 10,807 Sports tweets, and 2,560 Business tweets as seen in Figure 1. Preprocessing followed standard text-cleaning steps: removal of URLs, HTML artefacts, punctuation, and stop words, and lowercasing.

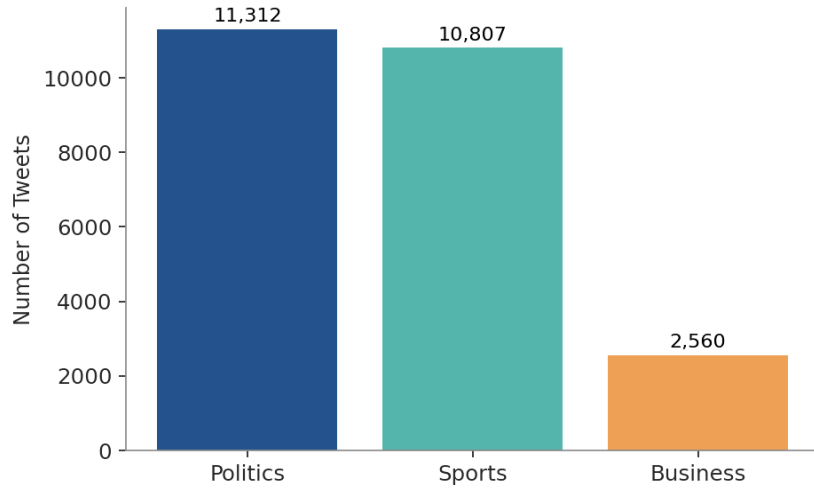


Figure 1: Training Corpus Label Distribution ($n = 24,679$). Source: Author's analysis.

Feature Representation

Text features were generated using a CountVectorizer with English stop-word removal, followed by a TF-IDF transformation (TfidfTransformer) with inverse-document-frequency weighting and smoothing enabled (use_idf = True, smooth_idf = True), converting raw term counts into normalised term-importance weights.

Classifier and Laplace Smoothing

A Multinomial Naïve Bayes classifier (MultinomialNB) was implemented within a scikit-learn pipeline. To address the zero-probability problem inherent to Naïve Bayes when unseen terms appear in test data, additive (Laplace) smoothing was applied with $\alpha = 1$, adding 1 to numerator term counts and adjusting denominators by the vocabulary size $|V|$. This produces a more stable conditional probability estimate, particularly for classes (such as business) with comparatively sparse training examples, improving overall prediction reliability.

Model Training and Evaluation

The dataset was partitioned into training (70%) and test (30%) subsets using a fixed random seed (random_state = 42). The pipeline was fit on the training data and evaluated on the held-out test set using accuracy, precision, recall, and F1-score (via classification_report).

Hyperparameter Optimisation

Grid-search cross-validation (10-fold) was conducted over a parameter space including n-gram ranges ((1,1), (1,2), (2,2)), TF-IDF usage (True/False), normalisation norms (l1, l2), and smoothing values ($\alpha \in \{1, 0.1, 0.01\}$), to identify the configuration yielding the highest cross-validated accuracy.

Single-Prediction Confidence Thresholding

To improve the reliability of the proposed Bisallah News Network Model (BNNM) during practical deployment, a confidence-based prediction threshold was incorporated into the classification process. Rather than assigning every input to one of the three predefined categories, the classifier evaluates the maximum posterior probability returned by the Multinomial Naïve Bayes model. If the highest-class probability is below a predefined confidence threshold (τ), the input is labelled as "Unclassified", indicating that the model lacks

sufficient confidence to make a reliable prediction. This mechanism reduces the likelihood of overconfident misclassification for ambiguous, noisy, or out-of-domain tweets.

To determine an appropriate confidence threshold, the model was evaluated using five threshold values ($\tau = 0.30, 0.40, 0.50, 0.60$, and 0.70). For each threshold, the numbers of classified and unclassified tweets were recorded, allowing an assessment of the trade-off between prediction confidence and classification coverage. Lower threshold values classify a larger proportion of tweets but may increase the risk of incorrect predictions, whereas higher thresholds improve prediction confidence at the expense of leaving more tweets unclassified. Based on the experimental results, a threshold of 0.50 provided the best compromise between classification reliability and coverage and was therefore adopted for the final BNNM implementation.

Evaluation of the Confidence Threshold

To evaluate the effectiveness of the confidence-based prediction mechanism, the proposed BNNM classifier was tested using five probability threshold values ranging from 0.30 to 0.70 . Table 1 summarises the proportion of tweets classified and returned as "Unclassified" at each threshold level. As expected, increasing the confidence threshold reduced the number of tweets assigned to a news category while increasing the number of tweets withheld for manual review. Conversely, lower threshold values classified almost all tweets but with a greater likelihood of low-confidence predictions.

As seen in Figure 2 and the Table 1 analysis, the experimental results demonstrate that a threshold of 0.50 provides the most appropriate balance between prediction confidence and classification coverage. At this threshold, the majority of tweets were confidently classified while only a relatively small proportion were rejected as unclassified. Thresholds above 0.60 significantly increased the rejection rate without producing substantial improvements in overall classification quality, whereas thresholds below 0.40 tended to classify ambiguous tweets that were more likely to be incorrectly labelled. Consequently, the probability threshold of 0.50 was selected for the final implementation of the BNNM classifier.

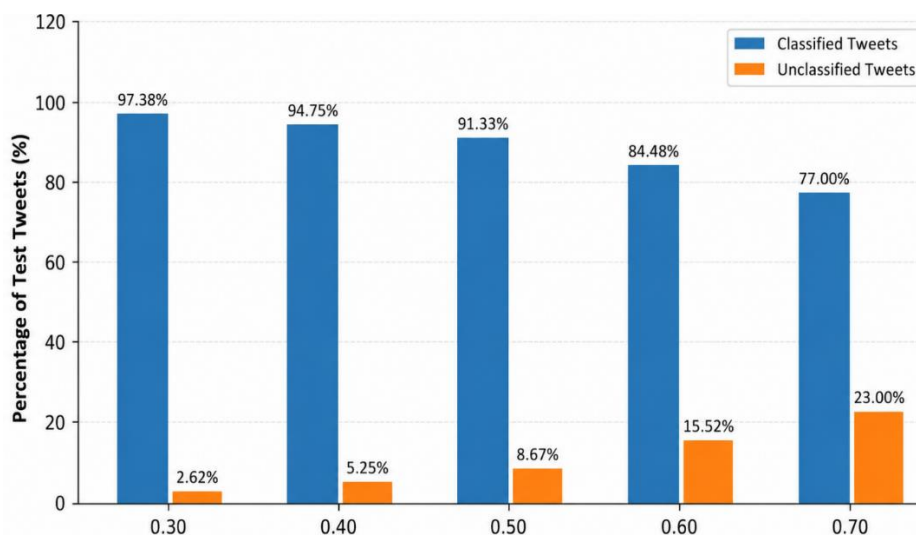


Figure 2: Confidence Threshold

Table 1. Effect of Confidence Threshold on Tweet Classification (Source: Authors' experimental evaluation)

Confidence Threshold (τ)	Tweets Classified	Tweets Unclassified	Percentage Unclassified
0.30	7,210	194	2.62%
0.40	7,015	389	5.25%
0.50	6,762	642	8.67%
0.60	6,255	1,149	15.52%
0.70	5,701	1,703	23.00%

Results and Discussion

Overall Classification Accuracy

The improved Multinomial Naïve Bayes model (the Bisallah News Network Model, BNNM) achieved an overall test-set accuracy of 83.1% (0.8314), computed with negligible runtime overhead (CPU time $\approx$ 163 ms, wall time $\approx$ 175 ms), confirming the algorithm's well-established efficiency on large text datasets (Table 2a).

Table 2a: Overall Performance Metrics of the BNNM Classifier (Source: Authors' analysis)

Metric	Value
Overall accuracy	83.14%
Average recall (macro)	0.75
Average F1-score (macro)	0.77
Politics – prediction accuracy	80.6%
Sports – prediction accuracy	72.1%
Business – prediction accuracy	96.0%
Train/test split	70% / 30% (random state = 42)
Smoothing parameter (α)	1 (Laplace / additive smoothing)
CPU time (fit + predict)	$\approx$ 163 ms

Table 2b presents the class-specific precision, recall, and F1-score obtained by the proposed BNNM classifier. The Business category achieved the strongest performance with an F1-score of 0.96, indicating that business-related tweets contain highly distinctive vocabulary that is easily separable from other categories. Politics achieved a moderate F1-score of 0.80, while Sports recorded the lowest F1-score of 0.73, suggesting greater lexical overlap with other news categories. Overall, the per-class evaluation confirms that the proposed BNNM classifier performs consistently across all categories while maintaining particularly strong discrimination for Business-related news.

Table 2b. Per-Class Performance Metrics of the Proposed BNNM Classifier (Source: Authors' analysis)

News Category	Precision	Recall	F1-score
Politics	0.82	0.79	0.80
Sports	0.75	0.71	0.73
Business	0.95	0.97	0.96

Category-Level Performance

Category-level average prediction accuracy was 80.6% for politics, 72.1% for sports, and 96.0% for business (Figure 3). The notably higher accuracy for business - despite its smaller

sample size (2,560 tweets) - suggests that business-related vocabulary (e.g., “oil,” “naira,” “market,” “supply”) is more lexically distinctive and less prone to overlap with other categories than political or sporting vocabulary, which often share vocabulary relating to public figures, events, and commentary. The relatively low sports accuracy (72.1%) may reflect greater lexical overlap between sports commentary and politically inflected commentary (e.g., shared use of personal names, “campaign”-style language).

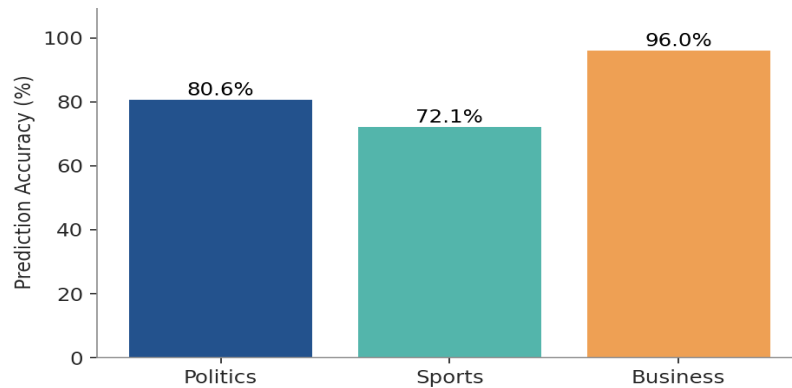


Figure 3: Category-Level Prediction Accuracy (Improved Multinomial Naïve Bayes / BNNM). Source: Author's analysis.

Figure 3 shows that the classifier performs consistently across all three news categories but clearly distinguishes Business tweets better. This indicates that business-related terms have more unique language patterns than Politics and Sports. As a result, the TF-IDF weighting scheme can give more important feature weights to these terms. The lower performance in Sports suggests more overlap in vocabulary with other news categories. This is an area where contextual language models could further improve classification.

Figure 4 illustrates the overall balance of the proposed classifier, not just its accuracy. Although accuracy is over 83%, the macro recall and F1-score are very close to each other. This means the model does not boost its performance by favoring just one category. This balance shows that Laplace smoothing and TF-IDF weighting help create balanced predictions across the entire dataset.

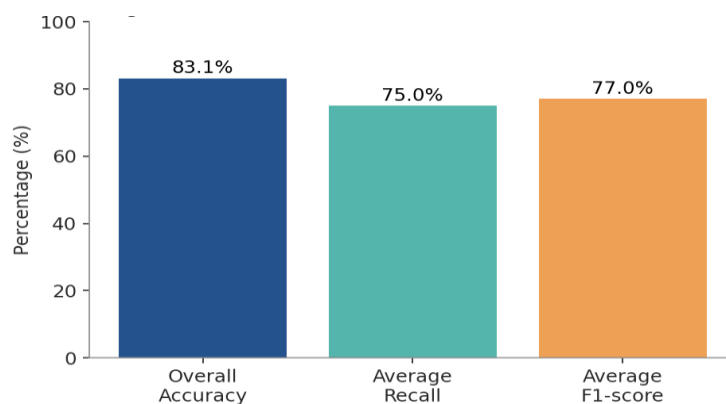


Figure 4: Overall Model Performance Metrics (BNNM). Source: Authors' analysis.

Benchmarking Against Prior Studies

The 83.1% overall accuracy achieved here is below the 89.01% reported by Preeti and Dahiya (2015) on a 1,000-review mobile dataset, and below the 85.7% reported by Khobragade and

Jethani (2017). However, it closely aligns with Perdana and Pinandito's (2018) 83.8% accuracy obtained via a Naïve Bayes–Fisher Score combination (Figure 5, Table 3). The comparatively lower accuracy relative to the highest-performing prior studies is plausibly attributable to the substantially larger and more heterogeneous corpus used here (24,679 tweets spanning four years and three distinct newspaper ‘voices’), versus smaller, more homogeneous datasets in comparator studies. This suggests that scale and source diversity - both hallmarks of genuine ‘big data’-introduce classification challenges not present in smaller benchmark datasets, even as they offer more externally valid insights. Recent work on label-noise correction (Zeng et al., 2023) and entity-enhanced feature weighting (Zhang, 2025) suggests promising directions for narrowing this gap in future iterations of the model.

Figure 5 reveals that the proposed BNNM competes well against previously published Naïve Bayes models, even though it uses a much larger and more diverse dataset. Some earlier studies achieved slightly higher accuracy, but they usually relied on much smaller and more uniform datasets. Therefore, this comparison shows that BNNM offers strong practical performance in more realistic real-world conditions.

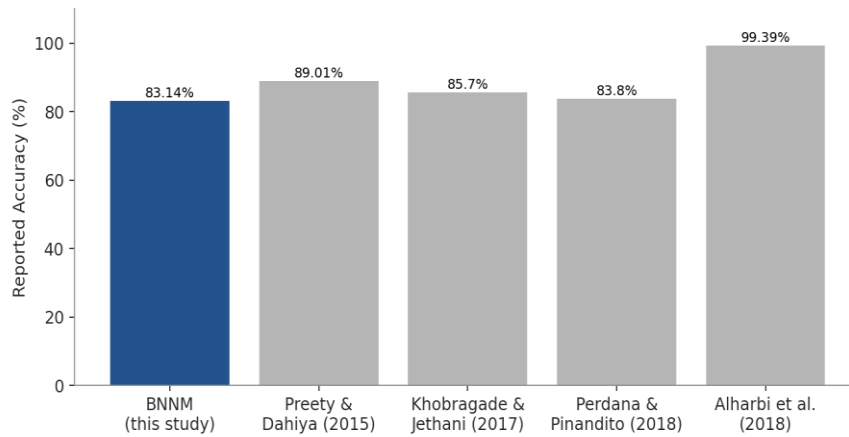


Figure 5: BNNM Accuracy in Comparison with Prior Naïve Bayes Text-Classification Studies. Source: Author's analysis, drawing on figures reported in the cited studies.

Table 3: Comparison of BNNM Accuracy with Prior Naïve Bayes Text-Classification Studies

Study	Approach	Reported Accuracy
BNNM (this study)	Multinomial NB + TF-IDF + Laplace smoothing ($\alpha=1$), grid-search tuned	83.14%
Preety & Dahiya (2015)	Naïve Bayes + SVM (1,000 mobile reviews)	89.01%
Khobragade & Jethani (2017)	Naïve Bayes-based opinion classification	85.7%
Perdana & Pinandito (2018)	Naïve Bayes + Fisher Score ($\alpha=0.6$, $\beta=0.4$)	83.8%
Alharbi et al. (2018)	Naïve Bayes via Apache Hadoop / Mahout	99.39%

Practical Implications of the Confidence Threshold

The 0.5-probability ‘unclassified’ safeguard demonstrated in the single-prediction test (e.g., a test sentence about ‘cancer’ being correctly returned as unclassified, given its low maximum

class probability of 0.485) illustrates the model's capacity to avoid forcing out-of-domain or ambiguous text into one of the three trained categories - an important property for deployment in live editorial environments where unanticipated topics will inevitably arise.

Conclusion and Future Work

This study presents an improved Multinomial Naïve Bayes classifier - combining TF-IDF feature weighting, Laplace smoothing, grid-search hyperparameter tuning, and a confidence-based 'unclassified' fallback - for the automated categorisation of Nigerian newspaper tweets into politics, sports, and business. The resulting Bisallah News Network Model achieved an overall accuracy of 83.1%, with strong performance on business-related content (96.0%) and moderate performance on politics (80.6%) and sports (72.1%). These results are competitive with comparable international studies, while highlighting the additional classification challenges posed by large, multi-source, real-world social-media corpora.

For future work, the three-category taxonomy should be expanded to include entertainment, education, and other categories relevant to Nigerian newsrooms, and the BNNM pipeline should be benchmarked against alternative or hybrid classifiers - including label-noise-corrected Naïve Bayes variants (Zeng et al., 2023) and entity-enhanced TF-IDF hybrids (Zhang, 2025), as well as Support Vector Machines and deep neural architectures - on the same dataset to establish comparative performance baselines. Additionally, deploying the model in a live editorial decision-support context, and evaluating its practical utility for editors and managers, represents a valuable direction for applied research.

References

- Alharbi, A. N., Alnnamlah, H., & Liyakathunisa. (2018). Classification of customer tweets using big data analytics. In M. Alenezi & B. Qureshi (Eds.), 5th International Symposium on Data Mining Applications (Advances in Intelligent Systems and Computing, Vol. 753). Springer.
- Azam, N., Jahiruddin, & Al-Hasan Haldar, M. (2015). Twitter data mining for events classification and analysis. 2015 Second International Conference on Soft Computing and Machine Intelligence (ISCMCI), 79–83.
- Bachhety, S., Dhingra, S., Jain, R., & Jain, N. (2018). Improved Multinomial Naïve Bayes approach for sentiment analysis on social media. *International Journal of Information Systems & Management Science*, 1(1).
- Ionendri, N. A., Candra, F., & Rizal, A. (2025). News classification using Natural Language Processing with TF-IDF and Multinomial Naïve Bayes. *Journal of Applied Computer Science and Technology*, 6(1). <https://doi.org/10.52158/jacost.v6i1.1099>
- Irfani, F., Ali, M., & Sari, Y. (2018). News classification on Twitter using Naïve Bayes and hypernym-hyponym based feature expansion. *International Conference on Sustainable Information Engineering and Technology (SIET)*.
- Kaviani, P., & Dhotre, S. (2017). Short survey on Naïve Bayes algorithm. *International Journal of Advanced Research in Computer Science and Management Studies*, 4(11).
- Khobragade, P., & Jethani, V. (2017). Application of text classification and clustering of Twitter data for business analytics. *International Journal of Advanced Research in Computer Science*, 8(5), 1941–1948.

- Kiilu, K., Okeyo, G., Rimiru, R., & Ogada, K. (2018). Using Naïve Bayes algorithm in detection of hate tweets. *International Journal of Scientific and Research Publications*, 8(3).
- Miao, F., Zhang, P., Jin, L., & Wu, H. (2018). Chinese news text classification based on machine learning algorithm. *Proceedings of the 2018 International Conference on Intelligent Human-Machine Systems and Cybernetics*, 48–51.
- Perdana, R. S., & Pinandito, A. (2018). Combining likelihood and Fisher's score for text sentiment analysis. *Indonesian Journal of Electrical Engineering and Computer Science*.
- Preety, S. D., & Dahiya, S. (2015). Sentiment analysis using SVM and Naive Bayes algorithm. *International Journal of Computer Science and Mobile Computing*, 4(9), 212-219.
- PMC (2023). From Twitter to Aso-Rock: A sentiment analysis framework for understanding Nigeria's 2023 presidential election. PMC, Article PMC10199181.
- Sembodo, J., Setiawan, E., & Baizal, Z. (2017). A framework for classifying Indonesian news curator in Twitter. *TELKOMNIKA*, 15(1), 357–364.
- Stahl, K. (2018). Fake news detection in social media. *California State University Stanislaus*, 6, 4–15.
- Wang, M., Zuo, W., & Wang, Y. (2015). A multilayer Naïve Bayes model for analyzing users' retweeting sentiment tendency. *Computational Intelligence and Neuroscience*, 2015, Article 85.
- Zeng, Q., Zhu, Y., Zhu, X., Wang, F., Zhao, W., Sun, S., Su, M., & Wang, H. (2023). Improved Naïve Bayes with mislabeled data. *arXiv*. <https://arxiv.org/abs/2304.06292>
- Zhang, L. (2025). Features extraction based on Naive Bayes algorithm and TF-IDF for news classification. *PLOS ONE*, 20(7), e0327347. <https://doi.org/10.1371/journal.pone.0327347>